

《心理科学进展》审稿意见与作者回应

题目：群际主观价值表征的神经计算机制及催产素调控

作者：张荷婧，陈奕铭，马焱娜

第一轮

审稿人意见：

本文围绕“群际偏差”这一经典的社会心理学议题，尝试从计算认知与神经机制的角度提出“内—外群体主观价值参照”的理论框架，进而引入催产素作为潜在调控变量，选题具有较强的理论意义与现实价值。作者在文中系统梳理了内群体偏好、群际差异加工及催产素调节群际关系的相关研究进展，文献覆盖面较广，能够体现对该领域主要理论脉络与代表性实证工作的把握。整体而言，文章具有较好的创新性，结构完整，但还可从清晰性与规范性等方面做进一步改进与完善。

意见 1：文中对核心模型的基本设定在不同位置存在表述不一致的问题，需要加以澄清与统一。在“研究构想”栏目自检报告第 5 点中，作者提出“社会群体参照点模型”用于描绘内—外群体主观价值表征差异；而在正文第 3 节“研究构想”开头段落中，又表述为“在这个模型中参照点上个体对于内群体和外群体利益的主观价值表征是等同的”。前后关于“差异”与“等同”的表述在逻辑上容易引起理解歧义。建议作者明确区分“参照点本身的定义”与“围绕参照点的价值权重偏移”，以避免误解该模型是在假定内外群体价值无差异，还是用于表示差异的来源与方向。

回应：感谢评审专家指出我们在稿件不同位置对核心模型表述可能引发歧义的问题。我们接受该建议，并已将“参照点”“内—外群体主观价值表征差异”与“围绕参照点的价值权重偏移”等概念进行了明确区分与统一表述。

修订后我们明确：本研究中的“社会群体参照点 φ ”描述了个体在主观最优的内—外群体分配方案时的群际赋权差异，它能够连续、精准且个性化地刻画群际决策中个体对内—外群体利益的相对重视程度及其偏向。具体来说，“社会群体参照点 φ ”是一个角度参数，表示在圆周方案集约束下使预测偏好 $R\cos(\theta - \varphi) = \beta_1 \times \cos\theta + \beta_2 \times \sin\theta$ 达到最大时的主观最优权衡方向，“社会群体参照点”是外群体（ β_2 ）相对内群体（ β_1 ）的赋权比例，由 (β_1, β_2) 的相位角定义， $\varphi = \arctan(\frac{\beta_2}{\beta_1})$ 。

因此，本研究可通过 β_1 、 β_2 、 φ 等指标来刻画个体对内—外群体主观价值表征差异的来源、方向与幅度。其中，“社会群体参照点” φ 反映个体在内群体与外群体利益之间的主观权衡偏向： φ 越接近 45° ，表明内外赋权越对称（ $\beta_1 \approx \beta_2$ ），此时个体对于内群体和外群体利益的主观价值表征是近似的；当 $\varphi < 45^\circ$ 表明个体更偏向内群体（ $\beta_1 > \beta_2$ ）， $\varphi > 45^\circ$ 表明个体更偏向外群体（ $\beta_1 < \beta_2$ ）。“社会群体价值距离”被定义为角距 $(\theta - \varphi)$ ，其表征当前的分配方案与个体主观上最优的内—外群体分配方案之间的偏离程度，它可作为后续神经表征分析中的参数指标。

我们对第 3 节“研究构想”开头等关键段落的措辞进行了重写，统一了模型命名，并在研究 2 的模型说明与图示注释中补充了对于“参照点”“围绕参照点的价值权重偏移”等关键概念的解释（相关修改已用蓝色字体标注）。

意见 2: 在国内外研究现状述评部分, 2.1 与 2.2 中关于群体偏差的社会心理学理论解释、神经网络模型与价值计算视角之间的关系, 多以并列方式呈现, 对“主观价值计算”在既有理论体系中的独特位置凸显稍显不足。建议可以更明确指出现有理论主要停留于态度、情绪或功能系统层面的局限, 并据此引出计算建模视角的必要性, 以增强综述的批判性与理论指向性。

回应: 感谢评审专家对综述结构提出的关键建议。我们根据建议对第 2 部分“国内外研究现状述评”的论述逻辑进行了调整, 以更突出“主观价值计算”在既有理论体系中的独特位置与必要性。具体修改包括:

1) 在 2.1 末尾扩写过渡段, 明确指出积极自我假设、社会认同理论与进化视角等解释主要回答“群体偏差为何发生”(身份动机、情绪态度或功能意义), 但较难进一步回答“个体在具体社会选择中如何将内外群体收益折算为可比较的主观价值”, 从而引出计算建模视角在机制刻画与量化预测上的优势。

2) 在 2.2 中增加对神经网络/功能系统证据的“贡献一局限”说明: 现有研究虽揭示了群体加工涉及的关键脑区与网络分工, 但仍缺少将这些神经证据转化为可计算、可估计、可比较的价值规则, 因此对个体差异与情境调节的预测力有限。基于此, 我们进一步强调计算建模能够以参数化方式刻画内外群体价值权重/参照结构, 并为后续神经机制检验提供直接的可检验指标。

3) 在 2.4 “现有研究存在的局限”中补充强调: 当前研究的共同缺口在于难以在统一框架下打通态度/情绪/功能系统解释与具体选择行为预测之间的链条; 因此引入主观价值计算模型是实现机制性解释、个体化刻画与干预评估的重要路径。

通过以上修改(相关修改已用蓝色字体标注), 我们将 2.1—2.2—2.4 由“并列罗列”调整为“层层递进”的论证结构: 从现象与社会心理解释出发, 过渡到神经加工证据, 再自然导向计算建模的必要性与创新点, 从而增强综述的批判性与理论指向性。

意见 3: 在研究设计之间的衔接上, 个别地方的逻辑过渡仍有进一步加强的空间。例如, 从研究 1 的人群分类(个体主义者、平等主义者、狭隘利他主义者)过渡到研究 2 时, 当前更多是经验性承接, 但尚未在理论层面明确说明为何该分类是后续模型建构与神经机制检验的“必要前提”。建议作者在研究 1 与研究 2 之间增加几句过渡性说明, 明确该分类在整体研究构想中的功能定位, 以增强研究路线的整体连贯性。

回应: 感谢审稿专家关于研究路线衔接的建议。我们接受该意见, 并在研究 1 与研究 2 的连接处、研究 2 和研究 3 的连接处补充了过渡性说明, 明确研究 1 的三类人群分类在整体研究构想中的功能定位。具体而言, 该分类并非仅用于描述性分组, 而是用于: ①在“内—外群体价值赋权”这一关键计算维度上对样本进行分层, 确保后续模型估计具有足够的参数变异与可辨识性; ②通过选择在人际层面利他性相对接近、但群际赋权模式显著不同的典型类型(尤其是平等主义与狭隘利他主义者), 在一定程度上控制一般利他倾向的混淆, 从而更聚焦检验“群际价值计算差异”本身; ③为研究 3 的神经机制检验提供明确的组间对照框架, 使我们能够以该分类作为关键组间变量, 结合研究 2 模型参数开展针对性的神经机制比较与假设检验。相应修改已体现在 3.1 末尾、3.2 开头与 3.3 开头(相关修改已用蓝色字体标注)。

意见 4: 研究 3 对神经分析方法的描述相对集中, 涵盖 ROI 分析、多体素模式分析、动态因果模型及图论分析等多种技术路线, 整体思路完整, 但在“为何需要同时采用多种分析方法”

这一问题上，当前主要以方法并列呈现为主。建议作者适度强化不同分析方法之间的分工逻辑与互补关系说明，以避免给读者造成方法堆叠的印象，更好地体现研究设计的针对性与必要性。

回应：感谢审稿专家对研究 3 分析策略的细致建议。我们接受该意见，并已在研究 3 的方法描述部分补充说明“为何需要同时采用多种分析方法”及其分工逻辑与互补关系，以避免呈现为方法并列堆叠。我们在修订稿中将研究 3 的神经分析明确组织为一条由局部到系统、由相关到机制的递进证据链，各方法分别对应不同层级的研究问题，且前一环节的结果为后一环节提供约束与依据（相关修改已用蓝色字体标注）：

1. ROI/全脑 GLM 参数调节分析用于回答“哪个脑区加工与计算变量相关的神经变化”。在一般线性模型框架下，以试次层面的关键计算量（如社会价值距离差）作为参数调节回归量，检验哪些脑区的 BOLD 信号与该变量呈系统性协变，从而定位候选关键区域。
2. 多体素模式分析（MVPA）用于回答“表征层面是否存在可区分差异”，以检验平等主义个体与狭隘民族主义个体在群际决策时的大脑活动模式是否存在稳定、可解码的差异（并可进一步定位包含组别差异信息的关键脑区/网络），从而在“活动强弱”之外提供更细粒度的表征证据。
3. 动态因果模型（DCM）用于回答“关键节点之间如何交互”。在激活分析和多体素分析确定的关键节点基础上，检验节点间的方向性有效连接及其是否受到任务条件/关键参数变量（如参照点）的调制，并通过模型比较对不同回路假设进行机制层面的区分。
4. 图论网络分析用于回答“系统层面的网络组织特征”。在功能连接网络上量化拓扑组织指标（如整合/分离特征、全局/局部效率、模块化、中心性等），以评估局部节点/回路差异是否在全脑网络层面呈现一致的组织学表现。

通过上述修订，我们在正文中明确了各方法的层级定位与递进关系：脑区（定位与表征）—回路—网络），并强调它们并非简单叠加，而是围绕同一核心问题从不同分析层次提供相互支撑的证据，从而更好体现研究设计的针对性与必要性。

意见 5：部分关键概念在全文中虽反复使用，但其操作性内涵尚不够明确，尤其是“主观价值表征”“社会群体参照点”“价值权重”“价值距离”等核心术语。目前这些概念主要通过分散的任务描述与模型公式加以呈现，尚缺乏相对凝练的概念界定。

回应：感谢审稿专家的建议，我们在最新的修订稿中进行了如下修改（相关修改已用蓝色字体标注），以期提升论文的严谨性与可读性：

（1）我们已对正文修订，使用精练的语言对全文相关术语进行了重新界定，并对这些核心概念加以区分，统一表述；

（2）在“研究设想”部分，新增了“核心概念与操作性定义”的表格，明确了每个术语的心理学含义、在本研究中的操作化定义，以及其是否由模型参数/公式直接表达。

第二轮

审稿人意见：

作者在本轮修订中对前一轮意见回应较为充分，尤其是通过补充学术术语界定与衔接段落，使群际参照点模型的理论定位和研究链条较初稿更为清晰。新增表格对关键概念进行了

集中说明，也有助于提升研究构想类文章应具备的可读性与规范性。从整体结构看，三个研究的递进关系比上一稿更明确，文本表达的连贯性也有所增强，但当前稿件仍存在需要进一步修改之处。

回应：非常感谢审稿专家在本轮审阅中对稿件修订工作的肯定与建设性建议。我们已根据您的意见对稿件进行进一步完善。以下为逐条回复。

意见 1：鉴于本文的核心创新点依赖于计算模型与参数化表征，公式书写的严谨程度直接关系到模型含义是否能够被准确理解。建议作者对正文、表格与图注中涉及参数的表达形式进行全面核查，确保符号使用前后一致，并在公式排版上采用更为规范的数学表达方式，以避免读者对关键表达产生歧义。例如，回复意见中的“ $R\cos\theta - \varphi$ ”、图 2 注释中“ $R\cos\theta - \varphi = \beta_1 \times \cos\theta + \beta_2 \times \sin\theta$ ”和“ (β_1, β_2) ”等。与此同时，文中涉及“相位角”或“参照点参数”的定义关系，建议作者在文本中以更清晰的方式说明其与 β 参数之间的对应逻辑（例如通过简要推导或明确的等价变换说明），从而增强模型表达的可验证性与理论透明度。

回应：感谢您对模型表达严谨性的提醒。我们完全认同公式与符号的一致性直接影响模型可理解性与可验证性。针对该问题，我们已对全文进行系统性核查与修订，具体包括：

1. 统一符号体系：对正文、表格与图注中涉及参数的写法进行了全面核对，确保同一概念仅使用同一符号、同一字体/下标形式；并纠正了不规范写法与笔误（将“ Φ ”替换成“ φ ”）
2. 规范公式排版：对所有公式采用 Latex 公式编辑器编辑（括号、下标、希腊字母、等号两侧空格等），避免因排版造成歧义。
3. 修改了文中模型部分的表达、删减冗余信息、保留和强调重点信息、补充详细的推导过程，提升模型的易读性。重点强调了模型及其推导过程，说明参数 β_1 、 β_2 与群际参照点 φ 的关系，关系，以提升理论透明度与可检验性，具体修改见 3.2 部分，附上修改后的文字：

“……在该任务中，被试需要对一系列社会资源分配方案进行评价，即对每种分配方案中的利益分配结果给出满意度评分（ Pr ），任务中分配方案的设计基于以笛卡尔坐标系原点为圆心的一组正圆轨迹（图 2a），其中横轴（X 轴）表示内群体的收益或损失，纵轴（Y 轴）表示外群体的收益或损失。借助这种设计，我们可以利用分配方案在坐标系中与坐标轴形成的角度（ θ ）来表示分配给内群体的利益和外群体的利益，根据个体对圆周上不同社会资源分配方案的满意度（图 2b），我们可以对每个被试的偏好进行数学建模（图 2c），具体的模型推导过程如下：

首先，假设个体对某一分配方案的主观价值（偏好）取决于该方案给予内群体和外群体收益的加权和，可表示为：

$$Pr = \beta_1 \cdot x + \beta_2 \cdot y$$

其中 Pr 表示主观偏好评分， x 和 y 分别表示该分配方案中给予内群体和外群体的客观收益， β_1 和 β_2 分别表示个体对内群体和外群体收益的主观权重系数。为方便建立模型，我们将分配方案在坐标系中的位置用极坐标表示， θ 为该方案点相对于 x 轴的极角（表示内外群体分配比例所对应的角度），方案点到原点的距离为常数 r （图 2a，其代表圆的半径，用以描述资源分配的总量），则可表示为：

$$Pr = \beta_1 \times r \cos\theta + \beta_2 \times r \sin\theta$$

考虑 r 为常数，简化后：

$$Pr = \beta_1 \times \cos\theta + \beta_2 \times \sin\theta$$

同样地，我们可以用极坐标形式表示个体的权重向量：设 φ 表示个体主观权重向量相对于 x 轴的角度（即“群际参照点”角度参数）， R 表示权重向量的模长（反映总体权重的尺度

大小)。则有：

$$\beta_1 = R \cos \varphi; \quad \beta_2 = R \sin \varphi$$

对 β_1 、 β_2 平方后求和，化简后可得：

$$\beta_1^2 + \beta_2^2 = R^2(\cos^2 \varphi + \sin^2 \varphi) = R^2;$$

故：

$$R = \sqrt{\beta_1^2 + \beta_2^2}$$

将 β_1 、 β_2 代入 Pr ，可得：

$$Pr = \beta_1 \times \cos \theta + \beta_2 \times \sin \theta = R(\cos \theta \cos \varphi + \sin \theta \sin \varphi)$$

根据余弦差角公式：

$$Pr = R(\cos \theta \cos \varphi + \sin \theta \sin \varphi) = R \cos(\theta - \varphi)$$

通过 β_1 、 β_2 两式相除可求群际参照点 ϕ ：

$$\frac{\beta_1}{\beta_2} = \frac{R \cos \varphi}{R \sin \varphi} = \tan \varphi;$$

$$\varphi = \arctan\left(\frac{\beta_2}{\beta_1}\right)$$

意见 2：在理论综述与研究动机部分，建议进一步消除少量语言层面的不规范表达，例如个别句子存在重复用字、标点使用不当或语句冗长的问题(2.2 第三、四段出现“发生系统性变化。。”“右侧眶额皮质（rOFC，承担价值计算功能）。（Liu et al., 2019）。”双句号错误；2.4 第二点中出现“这个问题的回答也许能为什么会有有群际偏差”语病等等)。这些问题虽不影响整体理解，但会降低研究构想类文章应具备的学术表达严谨性。建议作者在终稿阶段对全文进行一次系统性语言润色，重点检查长句拆分、逻辑连接词使用以及学术术语的统一性。

回应：感谢您指出语言层面的问题。我们已对您提到的具体位置逐一更正（如重复标点、语病与冗长句等），并在终稿阶段对全文进行多次系统性语言润色与一致性检查（修改部分用绿色标注），重点包括：

1. 标点与错别字/重复字（如双句号、重复用字等）；
2. 长句拆分与逻辑连接（提高句间衔接与可读性）；
3. 术语统一与表达规范。我们已对全文、表格与图注中涉及的关键术语进行了系统核查与统一，确保同一概念在不同位置的表述保持一致、规范且可追溯。与此同时，为使行文更简洁易读，并更符合中文学术写作中专业术语的常见构词与词组习惯，经慎重讨论，我们对模型名称进行了适度简化：将原“社会群体参照点模型”统一调整为“群际参照点模型”。后续文本中均采用该名称，并在首次出现处给出对应关系，以避免歧义并提升整体一致性与可读性，修改后语义更聚焦、术语更规范、表达更经济，且不改变模型的理论内涵。

我们相信该轮修订将进一步提升研究构想类文章应具备的表达严谨性与规范性

意见 3：在研究衔接方面，作者已尝试说明研究 1 的分类对于后续研究的重要性，但在部分段落中仍存在从分类结果直接推及理论结论的倾向。建议作者在措辞上进一步保持审慎，避免使用可能被理解为结论性或排他性推断的表达，而更多强调研究 1 在控制混淆、识别群体差异与为后续模型检验提供分组依据等方面的功能定位。这样的调整将更符合研究构想类文章“提出问题—建构模型—设计检验路径”的写作规范。

回应：非常感谢您对论证边界的提醒。我们同意在研究构想文本中应避免结论性或排他性措

辞。针对该问题，我们已对相关段落进行措辞层面的审慎调整：

1. 弱化研究 1 的结论化表述，改写为更符合构想论文定位的假设性/待检验性表达。具体修改见 3.1 预期结果部分，附上修改后的文字（修改后文字使用绿色标注）：

“我们预期研究 1 将检验一个关键假设：个体对内群体和外群体收益的主观价值赋权存在显著差异，并在价值计算中表现出明显的内群体偏好。同时，我们预期，这种在群际决策中的内群体利益偏好与个体在人际决策中的自利倾向并不完全重叠，这可能暗示存在一个独立于“人际决策情境下的亲社会—自利倾向”维度，即由内外群体情境引发的主观价值差异。如果数据支持这一假设，在考虑个体层面和群体层面的主观价值计算差异后，我们将对被试者进行聚类分析。通过聚类分析，我们计划识别出三类在人际价值偏好和群际价值偏好上表现出截然不同的个体：（1）个体主义者（Individualism）：其价值计算高度自我中心化。在人际互动中，他们优先最大化自身利益；在群际互动中，则倾向于将内群体利益等同于自身利益，从而表现出对内群体的显著偏好，对外群体缺乏积极关注。（2）平等主义者（Egalitarianism）：其价值计算强调公平与平衡。在人际交往中，他们表现出亲社会和利他倾向，在人际决策中注重公平；在群际互动中，他们赋予内群体与外群体利益近似相等的价值权重，从而在合作与资源分配上表现出跨群体的中立性与包容性。（3）狭隘利他主义者（Parochial altruism）：在人际交往中，他们表现出亲社会和利他倾向，在人际决策中注重公平；而在群际交往中，他们则表现出显著高估内群体利益、贬抑外群体利益，表现为强烈的内群体偏好与对外群体的排斥甚至敌意。”

2. 明确研究 1 的功能定位：更强调研究 1 在验证模型成立前提假设、识别个体差异、降低“亲社会—自利倾向”混淆、形成可解释的分组依据以及为后续参数化检验提供比较框架等方面的作用，而非作为理论结论本身。具体修改见 3.1 部分末尾和 3.2 开头，附上修改后的文字（修改后文字使用绿色标注）：

“如若此预期结果成立，后续研究将在研究 1 的基础上将重点关注“平等主义者”与“狭隘利他主义者”这两类被试，把他们作为组间因素纳入后续实验，以进一步检验超越个体利益的群际价值计算模式。之所以选取这两类极端被试，是因为他们在人际决策中的亲社会水平相近，可以控制了一般亲社会倾向的影响；同时，这两类被试在群际决策中的群际价值偏好迥异，便于我们针对“群际价值赋权差异”这一核心问题进行模型建构和检验。为此，我们计划从大样本中筛选出代表两种极端取向的被试（平等主义者与狭隘利他主义者，各 40 名，共 80 名）进入研究 2，以进一步建构“群际参照点模型”。

3.2 研究 2：构建“群际参照点模型”，以个体化描述内—外群体主观价值的计算过程

研究 2 以研究 1 为基础展开，旨在构建一种能够对内群体—外群体价值表征进行个体化描述的计算模型（图 2），用于量化经济利益、群体环境和社会规范等因素对内、外群体主观价值表征的影响。

研究 2 包含两项实验。实验一利用研究 1 中经过明确分类的被试，侧重于模型的构建；若实验一的模型拟合良好，实验二则随机招募普通被试，用于模型的验证。在研究 2 的实验一中，被试来源于研究 1 的大样本行为实验筛选。具体而言，我们基于研究 1 的结果预先筛选出两类人，他们在人际决策时均表现出亲社会倾向，而群际决策时在内群体—外群体利益主观偏好上截然不同——平等主义者（Egalitarianism，对内群体利益的主观价值与对外群体利益的主观价值相等的个体）和狭隘利他主义者（Parochial altruism，对内群体利益的主观价值远大于对外群体利益的主观价值的个体）。选取这两类极端被试有助于提高模型关键参数的可辨识性，从而增强模型拟合的稳定性与可靠性。”

3. 强化“提出问题—建模—检验路径”结构：在各个研究间的衔接段落处补充必要的过渡语句，突出后续研究将如何对关键机制进行模型化检验与神经表征验证。具体体现在各个研究的链接段落（修改后文字使用绿色标注）。

第三轮

审稿人意见：作者针对上一轮的意见进行了针对性回复。目前，我没有其他意见了。

编委复审意见：同意发表。